

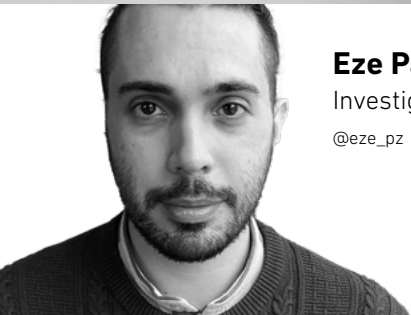


full de càlcul

Solidaritat amb cor, seny i mirada llarga: l'altruisme eficaç

vista prèvia >

La solidaritat és un dels valors més importants de cara a construir un món millor. Però la solidaritat, per millorar efectivament el món, ha d'estar ben orientada i dirigida. Cal decidir quins recursos destinem a resoldre quins problemes; quins mitjans i solucions posem en pràctica; i assegurar-nos que no deixem de tenir en compte el patiment de ningú. Aquestes són algunes de les qüestions a les quals s'adreça l'anomenat altruisme eficaç. Un corrent de pensament que no només reivindica el valor de la solidaritat, sinó que es proposa assegurar que l'acció solidària esdevingui una contribució eficaç i útil per a la millora de la vida al nostre món.



Eze Paez

Investigador Beatriu de Pinós al Law&Philosophy Group de la Universitat Pompeu Fabra

@eze_pz

Tenint en compte molts dels paràmetres de progrés més rellevants —com ara l'esperança i la qualitat de vida, la pobresa global o l'existència de violència—, el nostre món és avui millor que en qualsevol altre moment de la història.¹ Tot i així hi ha milions de persones que moren de malalties evitables. Les pitjors previsions del canvi climàtic auguren danys molts greus per a les nostres comunitats. I la proliferació d'armes nuclears i biològiques, o d'Intel·ligència Artificial, poden posar en perill l'existència mateixa de la humanitat.

Per a molts de nosaltres aquestes són qüestions de justícia de les quals s'haurien de fer càrrec els governs de les societats democràtiques enriquides. Tanmateix estem molt lluny de la realització dels ideals de justícia social, global i intergeneracional que es defensen des de posicions igualitàries. Àdhuc, cal reconèixer que en tota comunitat lliure hi ha d'haver un espai per a l'acció solidària de la societat civil. Allà on l'estat no hi arribi encara, i on no pugui o

no sigui convenient que hi arribi, ha de ser la ciutadania organitzada qui prengui la iniciativa.

On dirigir la nostra solidaritat? És indiferent o hi ha de causes millors que d'altres? La proposta de l'altruisme eficaç és que en l'acció solidària no és suficient posar-hi el cor, sinó que també tenim el deure de posar-hi el cap. Els nostres esforços per ajudar els altres han d'estar basats en raons i en evidències. Quan ho fem, descobrim que hi ha maneres millors i pitjors de ser solidari. Si ajudem de la manera més eficaç possible, aconseguirem que la vida de més persones sigui millor. En aquest text explicarem els diferents elements de la concepció de la solidaritat defensada des de l'altruisme eficaç.

Ajudar amb seny

Quan adoptem un objectiu, sigui quin sigui i en l'àmbit que sigui, la nostra primera tasca ha de ser identificar quines són les estratègies més eficients per assolir-lo. És a dir, les que ens permeten aconseguir en major mesura els nostres fins amb un cost

menor. Com a agents racionals cal que seguim l'estratègia que, d'acord amb les evidències disponibles, sembli la més adequada.

Un objectiu pot ser l'acció solidària. De fet molts creiem que, almenys en certes ocasions, tenim «l'obligació moral» d'ajudar els altres. El filòsof australià Peter Singer (1946) ho sol explicar amb el següent exemple:² imagineu que de camí a l'escola, la universitat o la feina passeu prop d'un estany poc profund. Un matí observeu que un nen ha caigut dins i s'està ofegant. Ficar-te dins l'estany i treure el nen seria fàcil, però acabaràs amb la roba tota molla i plena de fang. Si pretens anar a casa i canviar-te, quan tornis a l'estany ja hauràs fet tard. Una majoria de posicions ètiques conclourien que en aquest cas tenim el deure de salvar el nen. La intuïció de fons és que si podem fer un gran bé als altres sense cap cost per a nosaltres o amb un cost molt baix, llavors tenim l'obligació moral d'ajudar-los.

1 ROSER; ORTIZ-OSPINA; RITCHIE, «Life Expectancy»; ORTIZ-OSPINA; ROSER, «Global Health»; ROSER; ORTIZ-OSPINA, «Global Extreme Poverty»; ROSER; RITCHIE, «Homicides»; ROSER, «War and Peace».

2 SINGER, «The Drowning Child and the Expanding Circle».

Més enllà d'en quines circumstàncies estem obligats a ser solidaris, el cert és que totes les teories morals plausibles admeten que l'altruisme està moralment permès

Aquest exemple és força representatiu de moltes de les situacions reals en què un ciutadà «normal» dels països enriquits pot fer un bé als altres. La majoria no som rics, ni tenim responsabilitats polítiques ni formem part de grups de pressió. Però si ets ciutadà d'un país occidental, probablement guanyis més d'11.500€ l'any. Potser no et semblen ingressos gaire destacables, però en termes globals et situen dins del 10% més ric de la població. De fet, el 20% més pobre de la humanitat ha de subsistir amb poc més de 500€ anuals. Per tu, tenir un euro de menys pot no significar res. Pels més empobrits del planeta, tenir-ne un de més fa una diferència enorme.³ Així doncs, tenim un gran poder per fer el bé a baix cost.

És cert que en alguns casos el cost d'ajudar pot ser considerable. Quan en sorgeix, doncs, l'obligació? L'ètica no és pas una ciència exacta. És difícil determinar quin hauria de ser per cada persona l'esforç altruista mínim al qual estem obligats. Probable-

ment, però, el lllindar està per sobre del que fem avui dia.

Més enllà de en quines circumstàncies estem obligats a ser solidaris, el cert és que totes les teories morals plausibles admeten que l'altruisme està moralment permès. En tot cas, un cop ja hem escollit actuar de manera altruista, sorgeix el deure moral secundari de fer-ho de la manera més eficaç possible.⁴

A qui ajudar?

Gairebé ningú negaria que tots els humans són destinataris legítims de la nostra ajuda. Suposem que sabem que amb un cost assumible per nosaltres podríem aconseguir que algú salvi la seva vida o es curi d'una malaltia greu. Què importa pas el seu gènere, el color de la seva pell o inclús la seva nacionalitat? Res d'això és èticament rellevant. Què importa si viu lluny o a prop de nosaltres? Tampoc la distància en l'espai és en ella mateixa important a efectes morals.

Sembla difícil negar que tenim un deure general d'altruisme cap a tots els éssers humans que es troben en una situació de necessitat. Això és compatible amb creure que certes relacions especials ens donen raons addicionals per preferir ajudar uns éssers humans per sobre d'uns altres. Per exemple, relacions familiars, d'amistat o de concitadania. Més enllà d'aquestes raons especials, però, l'existència d'un deure general d'altruisme ens ha de comprometre a una actitud d'«imparcialitat front les causes». Cal escollir on invertim els nostres esforços solidaris atenent únicament a qüestions d'impacte: com podem fer el major bé al món, des d'un punt de vista imparcial.

Tanmateix, si acceptem un deure imparcial d'altruisme, no sembla possible en coherència restringir-lo només als éssers humans. No som els únics individus amb benestar propi, vulnerables o susceptibles de trobar-nos en necessitat. Això és comú a tots els éssers sintients, aquells amb capacitat per experimentar estats mentals de gaudi i patiment. El consens científic és que tots els animals vertebrats, i alguns d'invertebrats com els pops,

³ Fins a 100 vegades la utilitat que té per nosaltres, segons MACASKILL, *Doing Good Better. A Radical New Way to Make a Difference*.

⁴ PUMMER, «Whether and Where To Give».

Una causa altruista tindrà prioritat sobre altres en la mesura que resolde-la faci un bé major, s'hi destinin menys recursos i el guany marginal de destinar-los-hi sigui major

ho són, havent-hi gran incertesa sobre la resta d'invertebrats.⁵ No està justificat, doncs, excloure els animals de la nostra solidaritat.

Mirada llarga

Front el vici del «curtterminisme», l'altruisme eficaç defensa la posició del «llargterminisme». El més important en la nostra acció solidària és el seu impacte a llarg termini.⁶ De la mateixa manera que és irrellevant «on» existeix un individu que necessita ajuda, també és irrellevant «quan» existeix. La ubicació en el temps d'algú és tan poc important des d'un punt de vista ètic com el seu gènere, color de pell o espècie.

De fet, la majoria d'individus que existiran en algun moment, no se situen en l'actual present, o en un futur proper, sinó en el futur llunyà. Per tant ens hem de preocupar, sobretot, que els individus que viuran d'aquí a centenars o milers d'anys tinguin les millors vides possibles. Amb les nostres decisions egoistes o solidàries sobre canvi

climàtic o sobre altres riscos existencials podem condicionar, per bé o per mal, com seran les vides dels nostres descendents més llunyans.

La concepció de la solidaritat defensiva per l'altruisme eficaç és incompatible amb els prejudicis que tradicionalment ens han portat a excloure els altres de les nostres cures. Ens impel·leix a desprendre'ns del localisme, tant espacial com temporal, i que deliberem tenint en compte l'impacte de les nostres decisions en tots els que s'hi puguin veure afectats. Siguin qui siguin, visquin on visquin i existeixin quan existeixin.

Com prioritzar les causes?

Una solidaritat eficaç exigeix que identifiquem les causes més merixedores de la nostra atenció altruista. Es denomina «priorització de causes» a la tasca de mesurar-ne la importància relativa i, així, jerarquitzar-les. Els criteris usualment emprats per fer-ho són tres:⁷

1. La magnitud d'un problema: quant de bé faria resolde'l.

2. La desatenció d'un problema: la quantitat de recursos que ja s'hi destinen.

3. La solucionabilitat d'un problema: en quina mesura destinar-hi més recursos ajudaria a resolde'l.

Una causa altruista tindrà prioritat sobre altres en la mesura que resolde-la faci un bé major, s'hi destinin menys recursos i el guany marginal de destinar-los-hi sigui major.

Magnitud

La magnitud d'un problema és la mesura de com de millor seria el món en cas que es resolgués. És dolent emmalaltir de pneumònia i quedar-se al llit una setmana, però és encara pitjor contraure una malaltia crònica incapacitant. La pèrdua de qualitat de vida en el primer cas és temporal i relativament baixa. En el segon és permanent i molt elevada. Per mesurar aquestes pèrdues en qualitat de vida és usual en l'altruisme eficaç emprar la noció *qua-*

5 LOW, *The Cambridge Declaration on Consciousness*.

6 ORD, *The Precipice*.

7 Veure, pel que segueix, WIBLIN, «One approach to comparing global problems in terms of expected impact».

El problema de l'eficiència de les intervencions no és menor, ja que les millors maneres d'ajudar poden ser molt més eficaces que les que són merament bones

lity-adjusted life year (QALY) [any de vida ajustat per qualitat].⁸ Un QALY és equivalent a un any de vida en perfecte estat de salut. Fraccions de QALY impliquen un benefici de menor extensió temporal o menor nivell de salut. Així, suposem que si una persona no rep atenció mèdica perdrà dos anys de plena salut. Usant aquesta mesura, diríem que, en aquest exemple, la nostra ajuda tindria un impacte positiu de 2 QALYs. Imaginem que el resultat fos, en canvi, que perdés sis mesos de vida en perfecta salut. El nostre impacte seria de 0,5 QALYs. Per la mateixa raó, el benefici d'una atenció mèdica també seria de 0,5 QALYs si evitès que una persona perdés un any de vida de qualitat equivalent a la meitat d'un estat de salut perfecte.

Què contribuiria a un món millor, acabar amb la pobresa o reduir els riscos d'una pandèmia global greu? Potser aturar el canvi climàtic? És un càlcul difícil, perquè tindrem incerteses sobre a quantes persones afecta un problema, quin impacte tindria en

la seva vida resoldre'l o quines conseqüències indirectes comportaria solucionar-lo.

Desatenció

Suposem, d'altra banda, que resoldre un problema impediria la pèrdua de gran quantitat de QALYs. D'això no se segueix necessàriament que aquest problema mereixi més el nostre esforç solidari que una altra causa de magnitud menor. Si moltes persones han invertit molta feina i recursos en una causa és molt probable que ja hagin estat aprofitades les oportunitats d'impacte més importants.

A vegades una causa ha estat desatesa per bones raons, com la seva poca magnitud o solucionabilitat. D'altres cops, la desatenció pot venir d'un canvi en el sistema de valors dels agents solidaris. Fins fa poc, poques persones es preocupaven pels éssers humans en països llunyans, pel benestar animal o per les generacions futures. Si, segons els nostres valors, aquestes causes són d'una gran magnitud, podem dir que van ser desateses per raons errònies.

Solucionabilitat

El criteri de solucionabilitat ens convida a preguntar-nos quant progressaríem en una causa si li dediquéssim més recursos. Respecte a un mateix problema pot ser que hi hagi diferents intervencions amb un grau similar d'èxit. Quan això succeeixi cal identificar quina és l'acció més eficient. És a dir, quina esperablement resoldrà una fracció major del problema per cada unitat de recurs invertit.

El problema de l'eficiència de les intervencions no és menor, ja que les millors maneres d'ajudar poden ser molt més eficaces que les que són merament bones. Veiem-ne un parell d'exemples.⁹

Un problema en els països empobrits és l'alt absentisme escolar de les noies. Transferir 1.000 dòlars —uns 930 euros— directament a les estudiants augmenta l'assistència a l'escola en 0,2 anys de mitjana. Invertir la mateixa quantitat en uniformes gratuïts l'augmenta en 7,1 anys. Però fer servir aquests diners en un pro-

8 MACASKILL, *Doing Good Better: A Radical New Way to Make a Difference*.

9 Ibidem.

La desigual distribució de recursos i la inacció dels països enriquits fan que els governs d'aquests estats empobrits no tinguin la capacitat de fer front a aquests problemes

grama de desparasitació d'escolars l'augmenta en 139 anys. Així, millor intervenció altruista —el programa de desparasitació— és més de dinou vegades més eficient que l'anterior.

Un altre problema en els països empobrits és la salut. Fer servir 1.000 dòlars per la promoció de l'ús del preservatiu per prevenir el VIH dona un benefici esperable de 2 QALYs. Però finançar teràpies antiretrovirals evita la pèrdua de 5 QALYs. Si prenem en consideració altres malalties podrem observar que hi ha intervencions molt més eficients. Mil dòlars invertits en la distribució de xarxes antimosquits per prevenir la malària tenen un guany esperable de 10 QALYs.

Catàleg de problemes globals

Explicat d'aquesta forma, certament, potser costa conceptualitzar de forma clara quins poden ser els problemes o les causes que responen a aquesta prioritització entorn a l'altruisme eficaç. A continuació s'apunten alguns dels problemes globals que els defensors de l'altruisme eficaç consideren més mereixedors dels nostres esforços solidaris,

segons el marc magnitud-desatenció-solucionabilitat. Aquestes no són conclusions inapel·lables, sinó consideracions provisionals revisables a la llum de millors arguments i noves evidències.

La salut als països empobrits

Hi ha desenes de milions de persones que viuen en països empobrits, de les quals se'n podria evitar la mort o millorar la qualitat de vida, si tinguessin accés a una millor alimentació, a un entorn salubre i a una sanitat adequada. Malalties com la malària, la tuberculosi, el VIH o la diarrea causen greus problemes de salut. Minven la capacitat i les oportunitats per educar-se i formar-se, per pensar i per treballar. Acabar amb elles tindria, a més, altres efectes positius en aquestes societats, fent-les més pròsperes i estables.

La desigual distribució de recursos i la inacció dels països enriquits fan que els governs d'aquests estats empobrits no tinguin la capacitat de fer front a aquests problemes. Què podem fer al respecte com a societat

civil, a més de pressionar els nostres governants? Segons els càlculs d'organitzacions de l'àmbit de l'altruisme eficaç, les intervencions més eficients són les dirigides als països de l'Àfrica subsahariana. Consistirien en la distribució de medicaments contra la malària, de xarxes antimosquits i de suplementes vitamínics, així com programes de desparasitació.¹⁰

Per descomptat, donar diners per aquestes intervencions no és l'única contribució possible. Si tenim les habilitats i la motivació adients potser podem ajudar més eficaçment treballant per alguna ONG de gran impacte que operi sobre el terreny. També podem treballar per organitzacions d'avaluació, millorant els processos per determinar l'eficiència de diferents intervencions. Podria ser que el més convenient des d'un punt de vista solidari és que féssim recerca biomèdica sobre aquestes malalties.

Canvi climàtic: l'escenari extrem

Si no aconseguim reduir les emissions de gasos d'efecte hivernacle,

10 GIVEWELL, «Our Top Charities».

Si no aconseguim reduir les emissions de gasos d'efecte hivernacle, l'escenari més probable és que les temperatures mundials augmentin entre 2,6 °C i 4,8 °C al 2100

l'escenari més probable és que les temperatures mundials augmentin entre 2,6 °C i 4,8 °C en el 2100. Això seria terrible, però hi ha escenaris pitjors. Segons alguns càlculs, hi ha un 10% de probabilitats que l'augment de temperatures sigui superior a 4,8 °C.¹¹

Què succeiria en aquest escenari extrem? El desglaç de Groenlàndia. Un augment del nivell del mar de fins a dos metres que podria ocasionar 200 milions de refugiats. Onades de calor extrema. Sequeres i disminució de la producció agrícola. Disrupcions ecosistèmiques que portarien a una major desertificació i acidificació de l'aigua. També es podria provocar un cercle viciós que acceleri el procés d'escalfament si s'allibera a l'atmosfera el metà líquid del subsòl a causa del desglaç del *permafrost*.¹²

Per minimitzar aquests riscos podem treballar per reduir les emissions de gasos d'efecte hivernacle. Això inclou fer pressió perquè els

estats actuïn, desenvolupar tecnologia que generi menys emissions i prevenir o revertir la desforestació. També podem investigar en geoenginyeria, la disciplina que estudia com intervenir a gran escala per limitar els efectes del canvi climàtic.

Riscos de catàstrofe global derivats de l'avenç tecnològic

El progrés tecnològic és un dels factors que més ha contribuït a millorar la nostra vida. Malauradament algunes tecnologies actualment existents i d'altres que podríem desenvolupar generen riscos catastròfics. Són riscos de dany greu i irreversible a la qualitat de vida humana a escala global.

Els més rellevants es coneixen com a riscos existencials. Aquests «amenacen amb causar l'extinció prematura de la vida intel·ligent originada a la Terra o la destrucció permanent i dràstica del seu potencial per un desenvolupament futur desitjable».¹³

Un desastre nuclear global n'és un exemple. També ho seria una pandèmia molt més letal que la causada per la Covid-19, i que fos fruit, per exemple, d'un patògen artificial dissenyat pels éssers humans.

Una altra font de risc existencial és la possibilitat d'una superintel·ligència artificial hostil. Una superintel·ligència artificial —súper IA— és un sistema amb habilitats intel·lectuals molt superiors a les de qualsevol ésser humà. Segons Nick Bostrom (1973), un dels majors experts i divulgadors en aquest tema, el desenvolupament de les súper IA és possible i, fins i tot, probable que succeeixi en aquest segle.¹⁴ Certament, si la humanitat aconsegueix crear una agència superintel·ligent al nostre servei i alineada amb valors de compassió i justícia, la nostra civilització faria un gran salt qualitatiu endavant. Tanmateix, hi ha el risc que no puguem sotmetre la IA al nostre control i que sigui hostil cap als éssers humans i la resta d'éssers sintients. No es tractaria tant que les seves intencions fossin malèvoles, sinó que

11 DUDA i KOEHLER, «Climate change –extreme risks».

12 GIVEWELL, «Extreme risks from climate change».

13 BOSTROM, «Existential Risk Prevention as a Global Priority».

14 BOSTROM, *Superintelligence: Paths, Dangers, Strategies*.

Hi ha el risc que no puguem sotmetre la IA al nostre control i que sigui hostil cap als éssers humans i la resta d'éssers sintients

«la IA ni t'estima, ni t'odia, però està feta d'àtoms que pot fer servir per a una altra cosa».¹⁵

Cal treballar per impedir una IA hostil i altres riscos existencials. Com contribuir-hi de la manera més eficient depèn de les nostres habilitats i motivacions. Podem treballar en el desenvolupament tècnic de seguretat en IA, com fan al Machine Intelligence Research Institute [Institut de Recerca de la Intel·ligència de Màquines]. Podem fer recerca en riscos catastròfics, com al Future of Humanity Institute de la Oxford University [Institut del Futur de la Humanitat]. Podem pressionar per introduir aquests problemes en l'agenda política.

Benestar animal

Històricament molta gent ha considerat que els interessos dels animals no importen tant com els dels humans, o potser que no importen en absolut. Un nombre creixent d'autors sosté que aquesta actitud és fruit

d'un prejudici tan injustificat com el sexisme o el racisme i que s'ha anomenat «especisme».¹⁶

Cal que abordem la solidaritat, també, sense aquest punt cec moral. Cada any pateixen i moren en activitats econòmiques més d'un bilió d'animals.¹⁷ Això és 125 vegades la població humana. Gairebé sempre l'objectiu perseguit és comparativament trivial. El cas més destacat és el de la indústria alimentària, on es concentra la majoria d'animals explotats. Una alimentació totalment vegetal és perfectament saludable.¹⁸ Des d'un punt de vista imparcial els interessos en no patir dels animals haurien de tenir més pes que el plaer obtingut en assaborir-los.

Sembla que una de les estratègies més eficients per reduir el patiment animal consisteix en l'activisme de difusió de valors per canviar

els nostres actituds i hàbits.¹⁹ També ho és la pressió a institucions públiques i empreses per establir estàndards de benestar animal més exigents. Cal incloure entre les intervencions prometedores el desenvolupament d'aliments fets amb proteïna vegetal de nova generació o carn cultivada en laboratori. A més, com a consumidores, podem decidir no comprar productes d'origen animal. Com a ciutadanes, podem votar a qui prometi tractar aquest problema.

Juntament amb el problema del benestar dels animals domesticats trobem el del patiment dels animals salvatges. Hi ha més d'un trilió d'animals a la natura amb vides, probablement, d'intens patiment, no pas per l'acció humana, sinó per causes naturals.²⁰ No tenim encara els coneixements o la tecnologia per intervenir a gran escala en el seu benefici. Tanmateix, podem fer difusió del problema o investigar per obtenir els mitjans per ajudar-los.

¹⁵ YUDKOWSKY, «Artificial Intelligence as a Positive and Negative Factor in Global Risk».

¹⁶ SINGER, *Animal Liberation*.

¹⁷ MOOD; BROOK, «Estimating the number of farmed fish killed in global aquaculture each year»; FAO, «Statistics division: Production, live animals».

¹⁸ MELINA; CRAIG; LEVIN, «Position of the Academy of Nutrition and Dietetics: Vegetarian Diets».

¹⁹ ANIMAL CHARITY EVALUATORS, «Recommended Charities».

²⁰ TOMASIK, «How many wild animals are there»; FARIA, *Animal Ethics Goes Wild: The Problem of Wild Animal Suffering and Intervention in Nature*.

Sabem que treballem amb evidències incompletes i recursos limitats; cal, doncs, investigar com obtenir millors conclusions sobre l'avaluació i solució de problemes

Quant a accions solidàries d'abast més petit, podem dur a terme campanyes d'alimentació o vacunació. Exemples d'organitzacions que treballen en aquest àmbit són Rethink Priorities, Wild Animal Initiative o Animal Ethics.

Metacauses

Les metacauses són aquelles causes la resolució de les quals ens ajuda a identificar i solucionar millor les prioritats globals. N'identifiquem tres de principals. Sabem que treballem amb evidències incompletes i recursos limitats; cal, doncs, investigar com obtenir millors conclusions sobre l'avaluació i solució de problemes. Una segona metacausa és la difusió dels valors vinculats a l'altruisme eficaç, com ara la benevolència imparcial, la racionalitat i la presa de decisions sobre la base d'evidències. I finalment, moltes de les decisions importants que afecten al futur de la humanitat són preses per institucions. Desenvolupar mecanismes que ens permetin millorar la presa de decisions institucional té també una gran importància instrumental.

Conclusió

Aquesta és només una introducció a les idees de l'altruisme eficaç. Per aprofundir-hi és recomanable llegir els llibres de William MacAskill (1987), Peter Singer (1946) i Toby Ord (1979) que trobareu a la bibliografia. El web d'altruisme eficaç conté també un glossari d'idees clau.²¹ Moltes organitzacions que he esmentat realitzen la seva pròpia recerca i tenen canals propis de divulgació. Llegir els seus posts o escoltar els podcasts és una bona forma d'introduir-se en la discussió.

Cor, seny i mirada llarga. Expandir la nostra solidaritat fins a incloure tots els individus vulnerables. Rigor en la tasca de determinar com fer el major bé possible. Posar els ulls en el llarg termini i en aquells que es troben en l'avenir. Aquesta és la proposta de l'altruisme eficaç. ■

21 EFFECTIVE ALTRUISM, «Introduction to Effective Altruism».

■ Bibliografia

80,000 HOURS. «Our current view of the world's most pressing problems» [en línia]. A *80,000 Hours*, maig de 2020. Disponible a: <www.80000hours.org>.

ANIMAL CHARITY EVALUATORS. «Recommended Charities». A *Animal Charity Evaluators*, de novembre de 2020. Disponible a: <www.animalcharityevaluators.org>.

BOSTROM, Nick. «Existential Risk Prevention as a Global Priority» [en línia]. A *Global Policy*, núm. 4 (1), pàg. 15-31, 2013. Disponible a: <www.existential-risk.org>.

BOSTROM, Nick. *Superintelligence: Paths, Dangers, Strategies*. Oxford: Oxford University Press, 2014.

DUDA, Roman; KOEHLER, Arden. «Climate change (extreme risks)» [en línia]. A *80,000 Hours*, 2016. Disponible a: <www.80000hours.org>.

FAO. «Statistics division: Production, live animals» [en línia]. A *FAO*, 2018. Disponible a: <www.fao.org>.

FARIA, Catia. *Animal Ethics Goes Wild: The Problem of Wild Animal Suffering and Intervention in Nature*. Tesi doctoral, Universitat Pompeu Fabra, Barcelona, 2016.

GIVEWELL. «Extreme risks from climate change» [en línia]. A *GiveWell*, novembre de 2013. Disponible a: <www.givewell.org>.

GIVEWELL. «Our Top Charities» [en línia]. A *GiveWell*, novembre de 2020. Disponible a: <www.givewell.org>.

LEWIS, Gregory. «Reducing global catastrophic biological risks» [en línia]. A *80,000 Hours*, març de 2020. Disponible a: <www.80000hours.org>.

LOW, Philip. *The Cambridge Declaration on Consciousness* [en línia]. 2012. Disponible a: <www.fcmconference.org>.

MACASKILL, William. *Doing Good Better. A Radical New Way to Make a Difference*. Londres: Guardian Book, 2015.

MELINA, Vesanto; CRAIG, Winston; LEVIN, Susan. «Position of the Academy of Nutrition and Dietetics: Vegetarian Diets». A *Journal of the Academy of Nutrition and Dietetics*, núm. 116 (12), pàg. 1970-1980, 2016.

MOOD, Alison; BROOKE, Phil. «Estimating the number of farmed fish killed in global aquaculture each year» [en línia]. A *Fishcount*, 2012. Disponible a: <www.fishcount.org.uk>.

ORD, Toby. *The Precipice*. Londres: Bloomsbury, 2020.

ORTIZ-OSPINA, Esteban; ROSER, Max. «Global Health» [en línia]. A *Our World in Data*, 2016. Disponible a: <www.ourworldindata.org>.

PUMMER, Theron. «Whether and Where To Give». A *Philosophy and Public Affairs*, núm. 44 (1), pàg. 77-95, 2016.

ROSER, Max. «War and Peace» [en línia]. A *Our World in Data*, 2013. Disponible a: <www.ourworldindata.org>.

ROSER, Max; ORTIZ-OSPINA, Esteban. «Global Extreme Poverty» [en línia]. A *Our*

World in Data, 2013. Disponible a: <www.ourworldindata.org>.

ROSER, Max; ORTIZ-OSPINA, Esteban; RITCHIE, Hannah. «Life Expectancy» [en línia]. A *Our World in Data*, 2013. Disponible a: <www.ourworldindata.org>.

ROSER, Max; RITCHIE, Hannah. «Homicides» [en línia]. A *Our World in Data*, 2013. Disponible a: <www.ourworldindata.org>.

SINGER, Peter. *Animal Liberation*. Nova York: HarperCollins, 1975.

SINGER, Peter. «The Drowning Child and the Expanding Circle» [en línia]. A *New Internationalist*, abril de 1997. Disponible a: <www.utilitarian.net>.

TOMASIK, Brian. «How many wild animals are there» [en línia]. A *Reducing Suffering*, 2009. Disponible a: <www.reducing-suffering.org>.

WIBLIN, Robert. «Health in poor countries» [en línia]. A *80,000 Hours*, abril de 2016. Disponible a: <www.80000hours.org>.

WIBLIN, Robert. «One approach to comparing global problems in terms of expected impact» [en línia]. A *80,000 Hours*, octubre de 2019. Disponible a: <www.80000hours.org>.

YUDKOWSKY, Eliezer. «Artificial Intelligence as a Positive and Negative Factor in Global Risk». A BOSTROM, Nick; CIRKOVIC, Milan (eds.). *Global Catastrophic Risks*, pàg. 308-345. Nova York: Oxford University Press, 2008.